

Configuration Manual

MSc Research Project
Data Analytics

Emmanuella Anita Odiaka
Student ID: X23282819

School of Computing
National College of Ireland

Supervisor: Abid Yaqoob

National College of Ireland
MSc Project Submission Sheet
School of Computing



Student Name: Emmanuella Anita Odiaka
Student ID: X23282819
Programme: MSc in Data Analytics **Year:** 2024-2025
Module: MSc Research Practicum
Lecturer: Abid Yaqoob
Submission Due Date: 11/08/2025
Project Title: Enhancing Dublin Rental Price Prediction by Integrating Policy Insights via Large Language Models and Explainable AI
Word Count: 1272 **Page Count:** 13

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project.

ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature: Emmanuella Anita Odiaka

Date: 06/08/2025

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission, to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project, both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

Configuration Manual

Emmanuella Anita Odiaka
Student ID: x23282819

Introduction

This report gives a detailed breakdown to the configuration and Implementation of the project.

1 Environment

This section details the environmental and system requirement setup that was used in implementing the project.

A Systems Configuration

Table 1: Software Configuration

Requirement	Specification
Programing Language	Python Version 3
Tools	Google Collab, WORD Open AI LLM
Google Drive	Access to Drive for the Data
Operating System	macOS Sequoia 15.6

Table 2: Hardware Configuration

Requirement	Specification
Processor	2.3 GHz 8-Core Intel Core i9
Memory	32 GB 2400 MHz DDR4
Computational resources	Google Colab's CPU, 12.7 GB RAM

B Dataset

The rental data was taken from Kaggle as seen in Figure 1. It originally had 2,718 rows of daft.ie rental data. The data was downloaded in CSV format and then uploaded to Google drive.

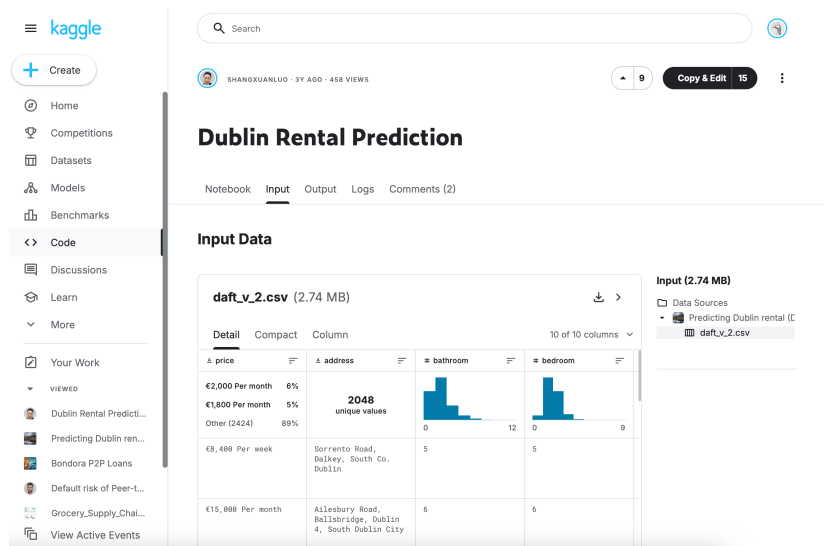


Figure 1: Rental Data Source

On the other hand, 14 government policy documents were coined from the internet using keyword searches like “Irish Rental Policy”. The report appendix shows the details of the documents used and the sources they were derived from.

C Google Colab Setup

The gathered data are stored in a Google Drive folder and accessed via Colab’s mount function as shown in Figures 2 and 3 respectfully.

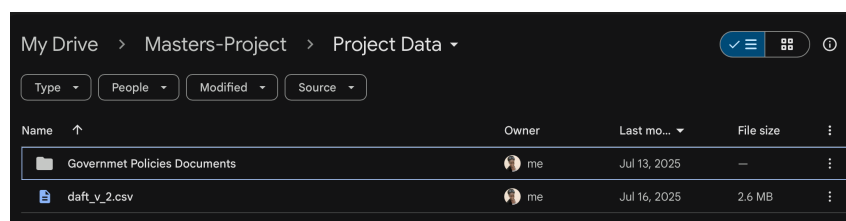


Figure 2: Data Location in Google Drive

The Mount function was used to build a connection between google drive and Colab to access the files as seen in Figure 3.

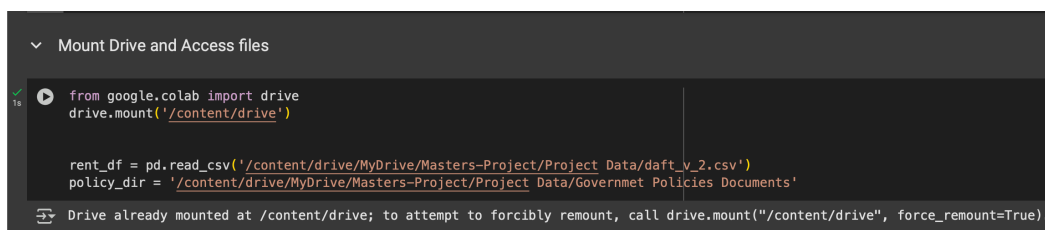


Figure 3: Loaded data using "Mount" from Colab

Figure 4 shows the selected language and processor engine used for the course of this project.

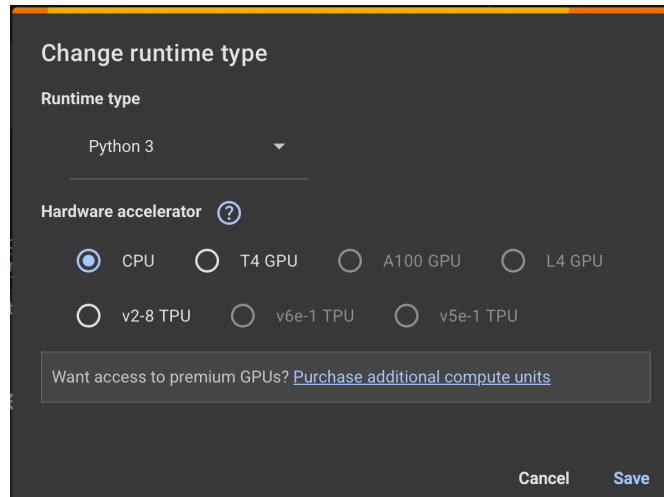


Figure 4: Colab Processor Engine

D Importing and Libraries

Figure 5 show the libraries imported and used for the implementation of the project.

```

import pandas as pd
import os
import numpy as np
import matplotlib.pyplot as plt
import seaborn as sns
import warnings # Ignore useless warnings (see SciPy issue #5998)
from sklearn.model_selection import train_test_split
from sklearn.ensemble import RandomForestRegressor
from sklearn.model_selection import train_test_split
from sklearn.preprocessing import StandardScaler, OneHotEncoder
from sklearn.model_selection import GridSearchCV
from sklearn.metrics import mean_absolute_error, mean_squared_error, r2_score

warnings.filterwarnings(action="ignore", message="^internal gelsd")

import PyPDF2
import openai
import json
from getpass import getpass
import shap
import lime
import lime.lime_tabular

```

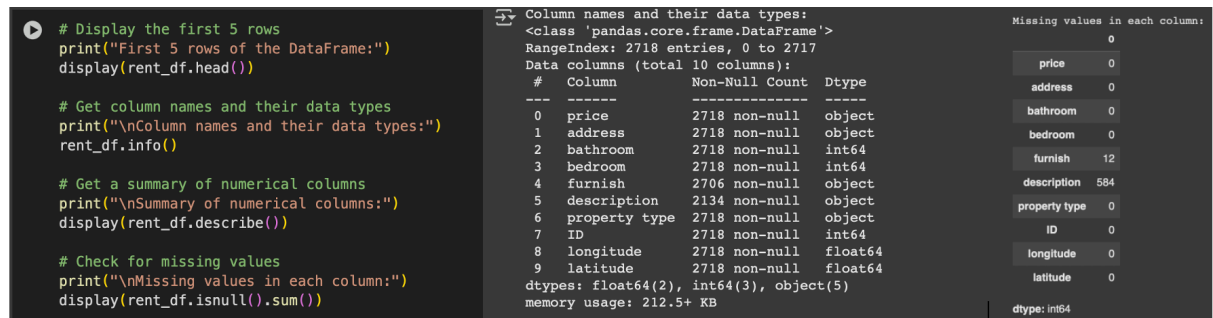
Figure 5: Libraries

2 Analysis and Evaluation of the Baseline model

This section will show the breakdown of the steps carried out to fully prepare and Analyse the baseline model.

A. Data Exploration and understanding

The process began by reading the basic information of the rental data and checking for null values and the variable types. This is seen in Figure 6.



B. Data Preparation

For this process, three tasks outlined below were performed after which the cleaned data was sent to the model building phase.

- **Data cleaning**

In this step, the missing values identified earlier were removed, price was converted to monthly values for consistency and the output was stored in a data frame for the Exploratory Data Analysis phase. This is seen in Figure 7.

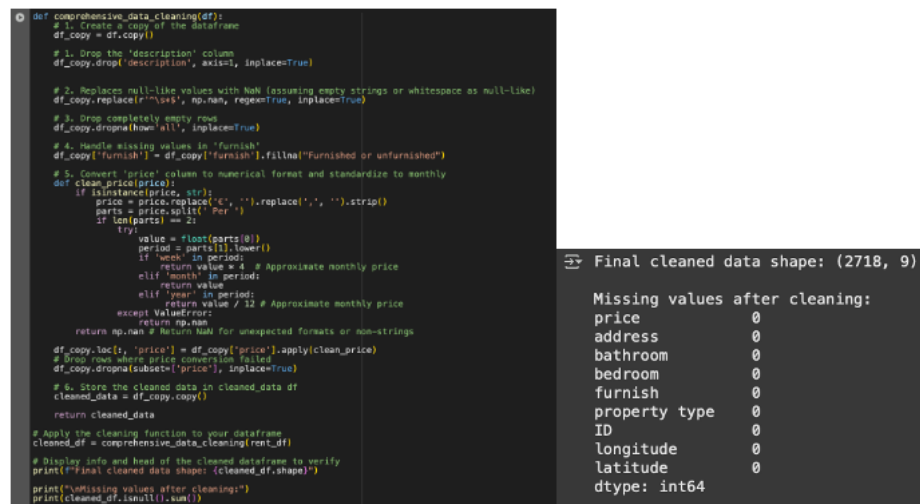


Figure 7: Data Cleaning Process

- **Exploratory Data Analysis**

His process involved performing a correlation analysis and viewing the relationship between the variables. A snippet of the full process is as seen in figure 8.

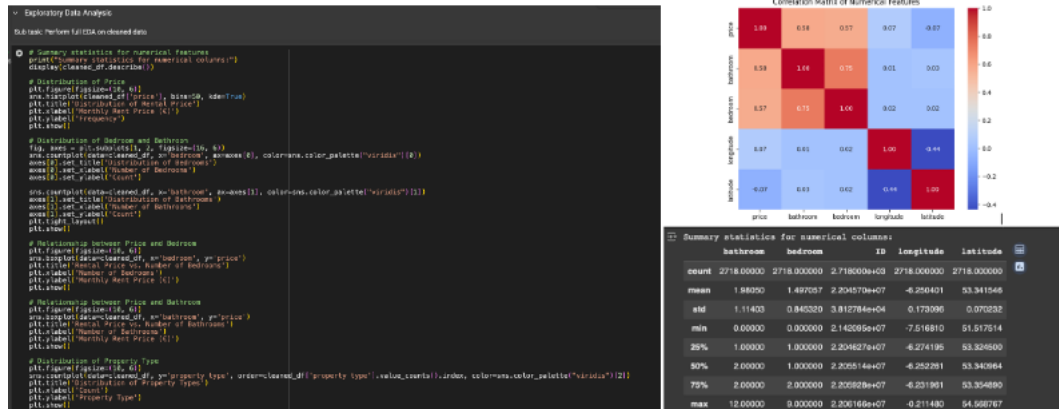


Figure 8: EDA Process and Results

• Outlier Handling

Here, outliers were detected and handled with the InterQuartile Range method, using the created 'price_per_bedroom' feature. This process is highlighted in Figure 9.

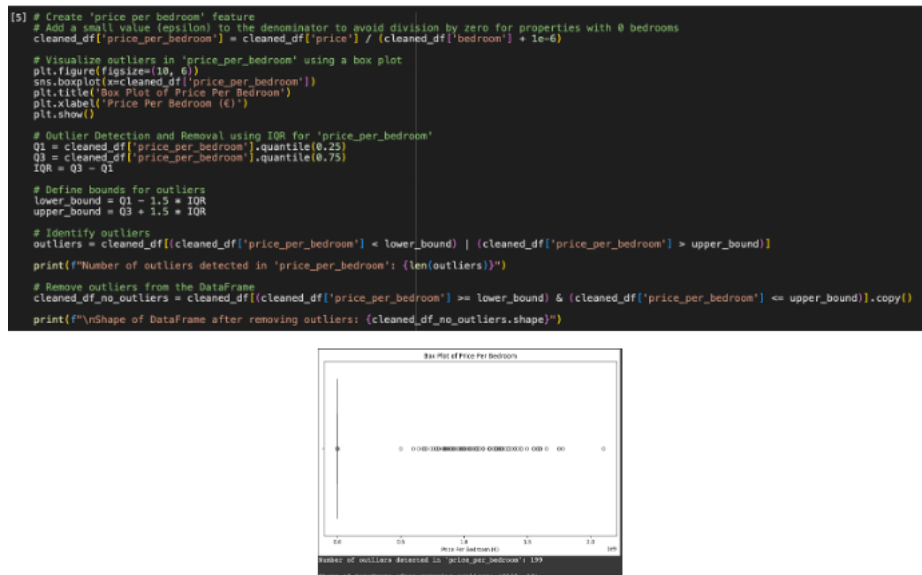


Figure 9: Outlier Handling Process and Result

• Feature Engineering and Onehot Encoding

Before sending the cleaned data for model building, the categorical variables were encoded and geographical features used in later analysis was engineered as seen in Figure 8.

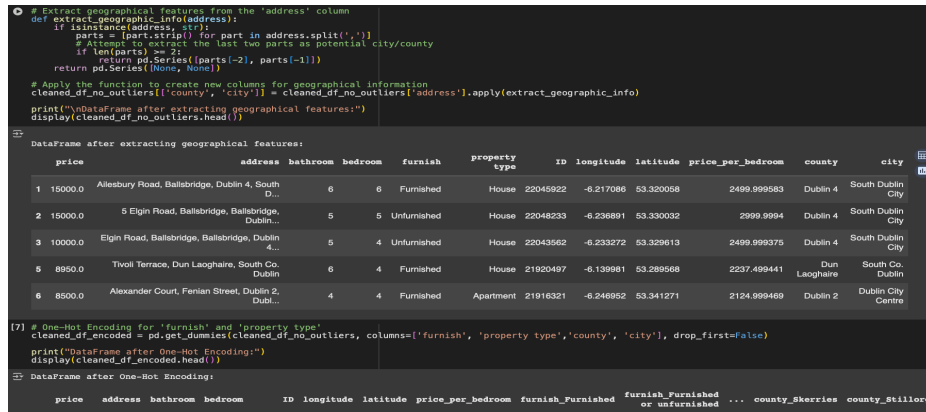


Figure 10: Feature Engineering & One-Hot Encoding

C. Model Building

The next step involved splitting the data into training and test sets with an 80:20 ratio and then sending the training data into the RandomForestRegressor model for training. GridSearch CV was used to find the best parameters for the regression model. All this is evidenced in Figure 11 below.

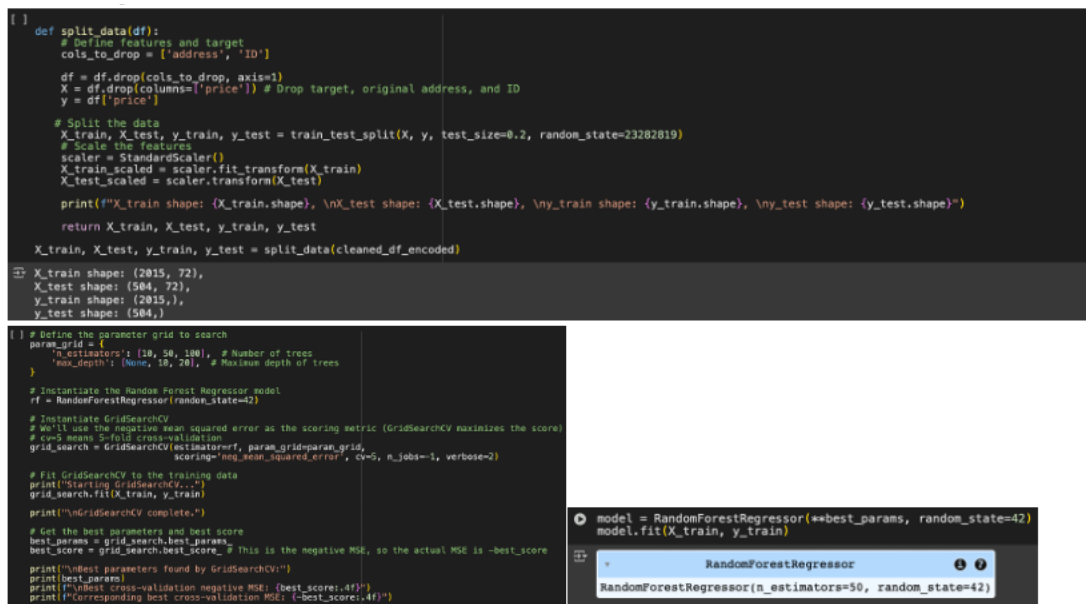


Figure 11: Baseline Modeling

D. Evaluation

Appropriate metrics were used to evaluate the performance of the model on the test data, after which feature importance analysis was carried out to see the top predictive features that most influenced the model, as shown in Figure 12.

```
[ ] # Make predictions on the test set
y_pred = model.predict(X_test)

# Calculate evaluation metrics
mae = mean_absolute_error(y_test, y_pred)
mse = mean_squared_error(y_test, y_pred)
rmse = np.sqrt(mse)
r2 = r2_score(y_test, y_pred)

# Print the evaluation metrics
print(f"Mean Absolute Error (MAE): {mae:.2f}")
print(f"Mean Squared Error (MSE): {mse:.2f}")
print(f"Root Mean Squared Error (RMSE): {rmse:.2f}")
print(f"R-squared (R2): {r2:.2f}")

# Get feature importances from the trained model
feature_importances = model.feature_importances_

if 'cleaned_df_encoded' in locals():
    df_for_features = cleaned_df_encoded.copy()
    target_column = 'price'
    X_for_features = df_for_features.drop(columns=[target_column, 'address', 'ID'])

    feature_importances_series = pd.Series(feature_importances, index=X_train.columns)

    # Get the top 10 most important features
    top10_features = feature_importances_series.nlargest(10)

    # Visualize the feature importances
    top10_features_sorted = top10_features.sort_values(ascending=False)

    plt.figure(figsize=(10, 7))
    sns.barplot(x=top10_features_sorted.values, y=top10_features_sorted.index, palette='viridis')
    plt.title('Top 10 Most Important Features (Baseline Model)')
    plt.xlabel('Importance')
    plt.ylabel('Feature')
    plt.show()

    print("Top 10 Most Important Features:")
    print(top10_features)
else:
    print("Error: cleaned_df_encoded DataFrame not found. Cannot determine feature names.")
```

Figure 12: Baseline Model Evaluation & Feature Importance

3 Prerequisites for the Policy Documents Processing

This section focuses on the steps taken to prepare the policy documents for the data extraction process.

A. Data preparation

Installs and file loading: First, the libraries needed for this section were installed and imported into the notebook as has been shown in Section 2D. After that, the policy files were loaded, and 2 samples were read to examine the structure and identify potential challenges that could be encountered in the extraction step. This is shown in Figure 13.

```
Load and explore the structure and content of the PDF Files

policy_files = os.listdir(policy_dir)
print("Files in the policy documents directory:")
for file in policy_files:
    print(file)

Files in the policy documents directory:
Planning and Development (Housing) and Residential Tenancies Act 2016.pdf
Residential Policy Covid.pdf
Affordable Housing Act 2021.pdf
Housing Policy: Actions to Deliver Change.pdf
Proposed Changes to Irish Residential Rent Controls.pdf
Revised Residential Tenancies and Valuation Act 2020.pdf
Residential Tenancies Bill 2021.pdf
Residential Tenancies (No. 2) Act 2021.pdf
Residential Tenancies Act 2021.pdf
RESIDENTIAL TENANCIES ACT 2004 REVISED Updated to 15 May 2025.pdf
housing-for-all-a-new-housing-plan-for-ireland-executive-summary.pdf
Cost rental housing.pdf
The Impact of Cost Rental Housing-5.pdf
Final Review RPZ and Policy Options for Rent Controls in PRS.pdf
Housing Assistance Payment (HAP).pdf

# Select a couple of files for exploration
selected_files = [
    'Planning and Development (Housing) and Residential Tenancies Act 2016.pdf',
    'Affordable Housing Act 2021.pdf'
]

print("\nExploring content of selected PDF files:")
for file_name in selected_files:
    file_path = os.path.join(policy_dir, file_name)
    print(f"--- Content of {file_name} ---")
    try:
        with open(file_path, 'rb') as f:
            reader = PyPdf2.PdfReader(f)
            num_pages = len(reader.pages)
            print(f"Number of pages: {num_pages}")
            # Read content from the first few pages to get an idea of the structure
            content = ""
            for i in range(min(num_pages, 3)): # Read first 3 pages
                page = reader.pages[i]
                content += page.extract_text() or "" # Use empty string for None
            # Print a truncated version of the content
            print(content[:1000] + "...") # Print first 1000 characters
    except FileNotFoundError:
        print(f"Error: File not found at {file_path}")
    except Exception as e:
        print(f"An error occurred while reading {file_name}: {e}")
```

Figure 13: Policy Documents Loading and Reading

OpenAI API Key Input: The API key used to access Open AI was prompted and stored as an environment variable and used to query the API as seen in Figure 14.

```

OpenAI API Key Input

# Use a prompt to enter the API key securely
# OPENAI_API_KEY:
# sk-proj-SALF4b-I_5bFSLXbyEwXiqm5bVlEjfj_9W80lseNCQBWCDEwxFPCVsZz-7Hm8xY-4Wbs2AI3hT3B1bkFJ2VJKH8zQbl

openai_api_key = getpass("Enter your OpenAI API key: ")
os.environ["OPENAI_API_KEY"] = openai_api_key

print(f"Policy document directory set to: {policy_dir}")

Enter your OpenAI API key: .....
Policy document directory set to: /content/drive/MyDrive/Masters-Project/Governmet Policies Documents

```

Figure 14: LLM API Accessibility

B. Extraction Logic

- **Initial Extraction Logic:** The process entailed defining a prompt template for the LLM and deriving the extracted information in chunks of 10000-character limit to send to the API and storing the resulting JSON output in a list that was later converted to a Pandas Dataframe. A snippet of this process is shown in Figure 15.

```

# Create a list of all PDF file paths within the policy directory
policy_files = [os.path.join(policy_dir, f) for f in os.listdir(policy_dir) if f.endswith('.pdf')]

extracted_data = []

# Define the prompt for the OpenAI LLM
prompt_template = """
Please extract the following information from the provided text about Irish housing and rental policies:
- Dates of policy implementation (e.g., "Effective from YYYY-MM-DD")
- Maximum allowed rent increase percentages (e.g., "Rent increases capped at X% per year")
- Specific criteria for rent reviews (e.g., "Rent review permitted every X months under conditions...")
- Geographical areas mentioned in relation to specific policies (e.g., "Rent Pressure Zones include...")
- Details about new housing supply targets or initiatives (e.g., "Target of X new homes by YYYY", "Initiatives include...")

Format the output as a JSON object with keys:
"source_document",
"implementation_date",
"rent_increase_cap",
"rent_review_criteria",
"geographical_areas",
"housing_supply_details". If information is not found, use "N/A".

Text to analyze:
{}
"""

# Iterate through each PDF file
for file_path in policy_files:
    print(f"Processing file: {os.path.basename(file_path)}")
    try:
        with open(file_path, 'rb') as f:
            reader = PyPDF2.PdfReader(f)
            text = ""
            for page_num in range(len(reader.pages)):
                page = reader.pages[page_num]
                text += page.extract_text() or ""

        # Implement a more robust chunking strategy for larger documents
        # Splitting by a larger delimiter or fixed size chunks
        # For simplicity, let's try splitting by a larger section marker if available, or fall back to character limit
        max_chunk_length = 10000 # Adjust based on model context limit and prompt length
        chunks = []
        if len(text) > max_chunk_length:
            # Simple character-based chunking for large texts
            for i in range(0, len(text), max_chunk_length):
                chunks.append(text[i:i+max_chunk_length])
        else:
            # Use the original simple split for smaller texts
            chunks = text.split('\n\n')
            if not chunks: # Handle cases where splitting by '\n\n' results in no chunks
                chunks = [text]

        for i, chunk in enumerate(chunks):
            if not chunk.strip(): # Skip empty chunks
                continue

            # Construct the prompt for the OpenAI LLM
            # Keep the prompt small to maximize text content in the chunk
            current_prompt = prompt_template.format(chunk)

            # Send the text chunk and the prompt to the OpenAI API
            try:
                response = openai.chat.completions.create(
                    model="gpt-3.5-turbo", # Or another suitable model
                    messages=[
                        {"role": "system", "content": "You are a helpful assistant that extracts specific information from text."},
                        {"role": "user", "content": current_prompt}
                    ],
                    response_format={"type": "json_object"}
                )
                llm_output = response.choices[0].message.content
            except Exception as e:
                print(f"Error processing chunk {i}: {e}")
    except Exception as e:
        print(f"Error processing file {file_path}: {e}")

```

Figure 15: Initial LLM Extraction Logic

The next step was to engineer features focused on policies related to Dublin, so the next step in the process was to study the information contained within the `geographical_areas` feature of the new dataset. This is shown in Figure 16.

```

Processing geographical_areas

print("Head of the 'geographical_areas' column:")
display(policy_df_original['geographical_areas'].head())

print("\nSample of 'geographical_areas' entries:")
# Display a few non-null unique values if there are a manageable number
non_null_areas = policy_df_original['geographical_areas'].dropna().unique()
if len(non_null_areas) < 50: # Display unique values if less than 50
    print(non_null_areas)
else: # Otherwise, display value counts for the most frequent entries
    print("Value counts for 'geographical_areas' (showing top 20 if many unique values):")
    print(policy_df_original['geographical_areas'].value_counts().head(20))

# Head of the 'geographical_areas' column:
Sample of 'geographical_areas' entries:
['Rent Pressure Zones', 'Rent Pressure Zones mentioned in the Act',
{'rent_pressure_zones': ['Cork City Council', 'Dublin City Council', 'Dún Laoghaire Rathdown Council'],
'rent_pressure_zones': 'Dublin, Kildare, Wicklow, Cork, Limerick, Westmeath, Wexham, Clare',
'rent_pressure_zones': 'N/A'},
{'Dublin', 'Limerick', 'Wexham Road Industrial Site'},
{'Affordable rental scheme', 'Valley', 'Error'},
{'Rent Pressure Zones', 'Rent Pressure Zones (RPS) extended nationally',
'Dún Laoghaire-Rathdown County Council'},
{'Rent pressure zones'},
{'Section 22', 'Amendment of section 22', 'Areas: Rent Pressure Zones'},
'Rent Pressure Zones include relevant areas within Local Government Act 2019',
{'Rent Pressure Zones', 'Ireland'},
{'Rent Pressure Zones': 'N/A'},
{'Rent Pressure Zones': 'Criteria: Annual rate of rent increase over 7% in at least 4 of 5 years',
'Limerick City and County Council', 'Kilkenny County Council', 'Wexham', 'Westport', 'Kildare City and County Council',
'Rent Pressure Zones include: Galway City West, Ballincollig-Carrigaline, Cork City Council,
Castlesher, Ballinacorney, Athlone (Kilcomman), Bligh-Carrigill, Carrigillow, Galway County Council,
Waterford City and County Council, Limerick City and County Council, Kilkenny County Council',
'Rent Pressure Zones', 'Dublin area'},
'Dublin for net household income below €46,000 per annum and everywhere else in the country for',
'Ireland', 'Dublin', 'Rest of the country'},
{'Kildare', 'Ballinacorney', 'Cork', 'Wexham', 'South Dublin'},
{'Type': 'Rent Pressure Zones', 'Designation': 'Areas experiencing consistent rental inflation',
'Dublin City'},
'Rent Pressure Zones including areas with consistent rental inflation in high-demand regions such as',
{'Rent Pressure Zones': 'Areas where price inflation fell relative to other regions'},
{'Rent Pressure Zones': 'RPS: Dublin, Other RPS: Non-RPS: Non-RPS areas'},
{'Areas mentioned': 'Rent Pressure Zones (RPS) in Ireland',
'Rent Pressure Zones include administrative areas and data in various local authorities and council',
'Rent Pressure Zones (RPS)'},
{'Rent Pressure Zones': 'Dublin City', '12.0% growth between 2016 and 2022', 'Dún Laoghaire-Rathdown',
'Dublin', 'Dublin', 'Greater Dublin Area'},
{'Dublin', 'GDA'},
{'Ireland'},
{'Germany'},
'Local Electoral Areas with Rent Pressure Zone system',
{'Pressure zones'},
{'RPS rent limits: Carlow County Council: N/A, Cavan County Council: N/A, Clare County Council: N/A',
'Dublin region'}]

```

Figure 16: Listing of Geographical Areas from Extraction

- **Refined Extraction Logic:** Based on how complex and varied the resulting data attributes from the data was, a new strategy with a refined logic was implemented to drive the extraction to focus specifically on Dublin area – a snippet of which is seen in Figure 17.

```

# Create a list of all PDF file paths within the policy directory
policy_dir = '/content/drive/MyDrive/Masters-Project/Government Policies Documents'
policy_files = [os.path.join(policy_dir, f) for f in os.listdir(policy_dir) if f.endswith('.pdf')]

extracted_data = []

# Refined prompt for the OpenAI LLM
refined_prompt_template = """
Please extract the following specific and location-aware information from the provided Irish housing
and rental policy text, focusing on details relevant to Dublin and its sub-areas:
- Dates of policy implementation (e.g., "Effective from YYYY-MM-DD")
- Maximum allowed rent increase percentages (e.g., "Rent increases capped at X% per year")
- Specific criteria for rent reviews (e.g., "Rent review permitted every X months under conditions...")
- Detailed Geographical Areas mentioned: List specific locations, focusing on Dublin postal codes,
sub-areas (like Ballsbridge, Dalkey, city council areas),
or broader regions if specific areas are not mentioned.
For each area, specify the <type> of area if mentioned
(e.g., "Rent Pressure Zone", "County", "City Council Area").
Explicitly state if a policy is noted as applying specifically to Dublin or parts of Dublin.
- Details about new housing supply targets or initiatives, noting if they are specific to Dublin.

Format the output as a JSON object with keys:
"source_document",
"implementation_date",
"rent_increase_cap",
"rent_review_criteria",
"geographical_areas_detailed",
"housing_supply_details".
For "geographical_areas_detailed", provide a list of objects, each with keys "area_name" and "area_type"
(use "N/A" if type is not specified). If information is not found or not relevant to Dublin,
use "N/A" for the respective key or an empty list for "geographical_areas_detailed".

Text to analyze:
{text}

extracted_data_refined = []

# Iterate through each PDF file (re-using policy_files list from previous steps)
for file_path in policy_files:
    print(f"Processing file: {os.path.basename(file_path)}")
    try:
        with open(file_path, 'rb') as f:
            reader = PyPDF2.PdfReader(f)
            text = ""
            for page_num in range(len(reader.pages)):
                page = reader.pages[page_num]
                text += page.extract_text() or ""

            # Implement a more robust chunking strategy for larger documents
            max_chunk_length = 10000 # Adjust based on model context limit and prompt length
            chunks = []
            if len(text) > max_chunk_length:
                # Simple character-based chunking for large texts

```

Figure 17: Refined LLM Extraction Logic

C. Hybrid Dataset creation

The steps followed to generate the hybrid dataset started from identifying and aggregating relevant key words or features from the extracted data and then linking them to individual rental properties data to generate the hybrid dataset. These steps are shown in Figures 18 – 20

Aggregate policy features creation is on the presence of specific keywords like rent increase, review criteria and housing supply details as seen in figure 18. Here, the merging of the original cleaned baseline dataset and the aggregated features was carried out.

```

# Let's create features based on the presence of specific keywords or non-'N/A' values
# We'll count how many entries in policy_df (which came from chunks of documents)

# Count non-'N/A' entries in relevant columns as a proxy for policy mentions
num_rent_increase_mentions = policy_df[policy_df['rent_increase_cap'] != 'N/A'].shape[0]
num_rent_review_mentions = policy_df[policy_df['rent_review_criteria'] != 'N/A'].shape[0]
num_housing_supply_mentions = policy_df[policy_df['housing_supply_details'] != 'N/A'].shape[0]

print(f"Number of policy entries mentioning rent increase caps: {num_rent_increase_mentions}")
print(f"Number of policy entries mentioning rent review criteria: {num_rent_review_mentions}")
print(f"Number of policy entries mentioning housing supply details: {num_housing_supply_mentions}")

# Create new columns in the original DataFrame with these aggregate counts.
df_hybrid = cleaned_df_encoded.copy() # Create a copy to avoid modifying the original df directly
df_hybrid['policy_rent_increase_mentions'] = num_rent_increase_mentions
df_hybrid['policy_rent_review_mentions'] = num_rent_review_mentions
df_hybrid['policy_housing_supply_mentions'] = num_housing_supply_mentions

print("\nDataFrame with added aggregate policy features:")
display(df_hybrid.head())

```

Number of policy entries mentioning rent increase caps: 180
 Number of policy entries mentioning rent review criteria: 180
 Number of policy entries mentioning housing supply details: 180
 DataFrame with added aggregate policy features:

Figure 18: Aggregated Hybrid Dataset Creation

Geographical areas features were standardised to normal location attributes like “Dublin”, “Lucan”, “Killiney”, etc as seen in Figure 16, which was then used to filter out Dublin specific policies from the refined dataframe as seen in Figure 19.

```

policy_df_refined = policy_df.copy()
def extract_and_standardize_locations(area_list):
    locations = []
    if isinstance(area_list, list):
        for area_info in area_list:
            if isinstance(area_info, dict) and 'area_name' in area_info:
                area_name = area_info['area_name']
                # Simple standardization: convert to lower case and remove leading/trailing whitespace
                standardized_name = area_name.lower().strip()
                # Further standardization involve mapping variations (e.g., "Dublin City Council" to "Dublin")
                locations.append(standardized_name)
    return list(set(locations)) # Return unique standardized locations

policy_df_refined['standardized_locations'] = policy_df_refined['geographical_areas_detailed'].apply(extract_and_standardize_locations)

# Examine the extracted standardized locations
print("\nSample of standardized locations extracted from policy documents:")
display(policy_df_refined['standardized_locations'].head())

# Get unique values from the 'county' column
if 'county' in cleaned_df_no_outliers.columns:
    unique_counties = cleaned_df_no_outliers['county'].unique()
    unique_counties = [str(item).lower() for item in unique_counties if isinstance(item, str)]
    print("Unique values in the 'county' column of cleaned_df_no_outliers:")
    print(unique_counties)
else:
    print("Error: 'county' column not found in cleaned_df_no_outliers.")

print("\n" * 30)

# Get unique values from the 'city' column
if 'city' in cleaned_df_no_outliers.columns:
    unique_cities = cleaned_df_no_outliers['city'].unique()
    unique_cities = [str(item).lower() for item in unique_cities if isinstance(item, str)]
    print("Unique values in the 'city' column of cleaned_df_no_outliers:")
    print(unique_cities)
else:
    print("Error: 'city' column not found in cleaned_df_no_outliers.")

# Given the variability, a simple approach is to create features indicating
# the presence of any Dublin-related location mention in a policy entry.
def is_dublin_related(locations):
    if not locations:
        return False
    # Check if 'dublin' or any Dublin-specific keywords are in the standardized locations
    dublin_keywords = [] # Add other general keywords as needed
    # Add unique Dublin counties and cities from your rental data (in lowercase)
    # Ensure you only add relevant Dublin-area counties/cities if your unique lists contain non-Dublin areas
    dublin_keywords.extend(unique_counties)
    dublin_keywords.extend(unique_cities)
    # Convert to a set and back to a list to remove duplicates and keep it clean
    dublin_keywords = sorted(list(set(dublin_keywords)))
    # print(f"Generated Dublin keywords for filtering: {dublin_keywords}")
    for loc in locations:
        if any(keyword in loc for keyword in dublin_keywords):
            return True
    # Also check if area_type indicates a Dublin-related zone if available
    # (This would require modifying the extract_and_standardize_locations to return type as well)
    # For now, rely on name matching.
    return False

policy_df_refined['is_dublin_related_policy'] = policy_df_refined['standardized_locations'].apply(is_dublin_related)

# Count the number of policy entries (chunks) that are considered Dublin-related
num_dublin_related_policy_mentions = policy_df_refined['is_dublin_related_policy'].sum()

# Count mentions of specific policy types relevant to Dublin (e.g., RPZ)
# This requires examining the original extracted data before standardization if type info is needed
# Let's re-process the raw geographical_areas_detailed to look for RPZ type mentions
def count_rpz_mentions(area_list):
    count = 0
    if isinstance(area_list, list):
        for area_info in area_list:
            if isinstance(area_info, dict) and 'area_type' in area_info and 'Rent Pressure Zone' in area_info['area_type']:
                count += 1
            elif isinstance(area_info, dict) and 'area_name' in area_info and 'Rent Pressure Zone' in area_info['area_name']: # Sometimes type is in name
                count += 1
    return count

policy_df_refined['rpz_mentions_count'] = policy_df_refined['geographical_areas_detailed'].apply(count_rpz_mentions)
total_rpz_mentions_in_dublin_related_policies = policy_df_refined[policy_df_refined['is_dublin_related_policy']]['rpz_mentions_count'].sum()

print(f"\nNumber of policy entries identified as Dublin-related: {num_dublin_related_policy_mentions}")
print(f"\nTotal count of RPZ mentions within Dublin-related policy entries: {total_rpz_mentions_in_dublin_related_policies}")

```

Figure 19: Standardization and Filtering of Refined Dataset

Finally, these features were merged with the aggregated policy features and this formed the newly defined Hybrid Dataset as shown in Figure 20.

```
[ ] # Count the number of policy entries (chunks) that are considered Dublin-related
num_dublin_related_policy_mentions = policy_df_refined['is_dublin_related_policy'].sum()

# Count mentions of specific policy types relevant to Dublin (e.g., RPZ)
# This requires examining the original extracted data before standardization if type info is needed
# Let's re-process the raw geographical_areas_detailed to look for RPZ type mentions
def count_rpz_mentions(area_list):
    count = 0
    if isinstance(area_list, list):
        for area_info in area_list:
            if isinstance(area_info, dict) and 'area_type' in area_info and 'Rent Pressure Zone' in area_info['area_type']:
                count += 1
            elif isinstance(area_info, dict) and 'area_name' in area_info and 'Rent Pressure Zone' in area_info['area_name']: # Sometimes type is in name
                count += 1
    return count

policy_df_refined['rpz_mentions_count'] = policy_df_refined['geographical_areas_detailed'].apply(count_rpz_mentions)
total_rpz_mentions_in_dublin_related_policies = policy_df_refined[policy_df_refined['is_dublin_related_policy']]['rpz_mentions_count'].sum()

print(f"\nNumber of policy entries identified as Dublin-related: {num_dublin_related_policy_mentions}")
print(f"Total count of RPZ mentions within Dublin-related policy entries: {total_rpz_mentions_in_dublin_related_policies}")
```

Figure 20: Final Hybrid Dataset

4 Policy Driven Implementation

A. Data Splitting

Following the derivation of the Hybrid dataset, the following steps include splitting of the Hybrid dataset into training and testing subsets using the split function earlier defined as shown in Figure 21

```
X_hybrid_refined, X_hybrid_refined_test, y_hybrid_refined, y_hybrid_refined_test = split_data(df_hybrid_refined)
print("\nPrepared Refined Hybrid Feature set (X_hybrid_refined) head:")
display(X_hybrid_refined.head())
print("\nPrepared Refined Hybrid Target variable (y_hybrid_refined) head:")
display(y_hybrid_refined.head())
print("\nInfo of the Prepared Refined Hybrid Feature set (X_hybrid_refined):")
X_hybrid_refined.info()

X_train shape: (2015, 77),
X_test shape: (504, 77),
y_train shape: (2015,),
y_test shape: (504,)

Prepared Refined Hybrid Feature set (X_hybrid_refined) head:
bathroom bedroom longitude latitude price_per_bedroom furnish_Furnished furnish_Unfurnished
```

Figure 21: Split Hybrid Data

B. Model Building and Training

The earlier defined regressor model along with the same parameters are also used to test the hybrid model as shown in Figure 22.

```
hybrid_model = RandomForestRegressor(**best_params, random_state=42)
hybrid_model.fit(X_hybrid_refined, y_hybrid_refined)
```

RandomForestRegressor

RandomForestRegressor(n_estimators=50, random_state=42)

Figure 22: Hybrid Model Training

5 Evaluation and Explainability

This section covers the evaluation performed on the dataset, and the tools used for the model's explainability.

D. Model Evaluation

The same feature importance analysis and metrics used to evaluate the baseline model, was used in this step as seen in Figure 23, along with a comparison with the baseline model's performance.

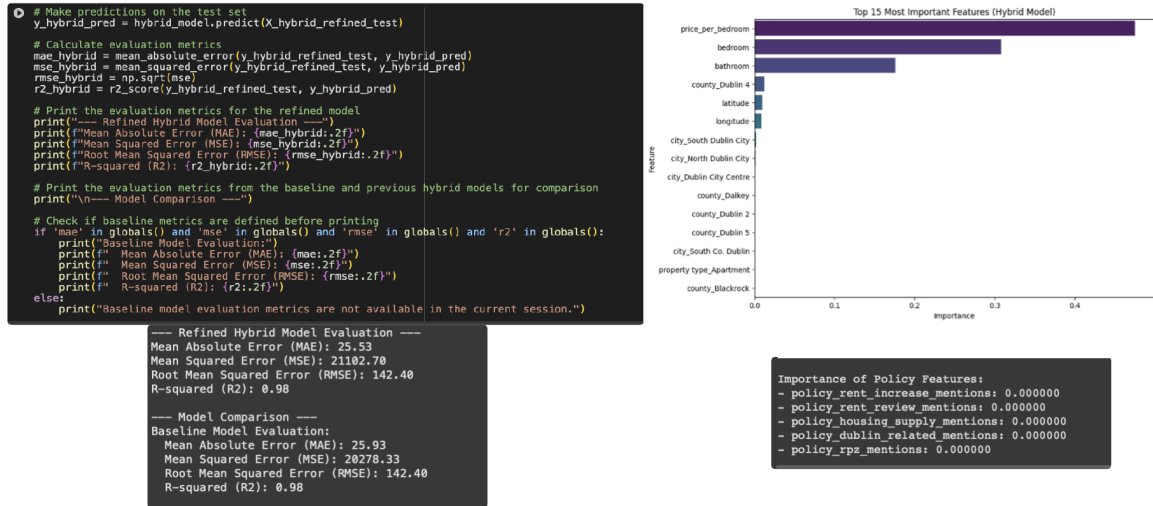


Figure 23: Hybrid Model Evaluation and Feature Importance

E. Model Explainability

The hybrid model was further evaluated using SHAP and LIME to explain the model's predictions. The implementation and results are as seen in figures 24 and 25 respectively.

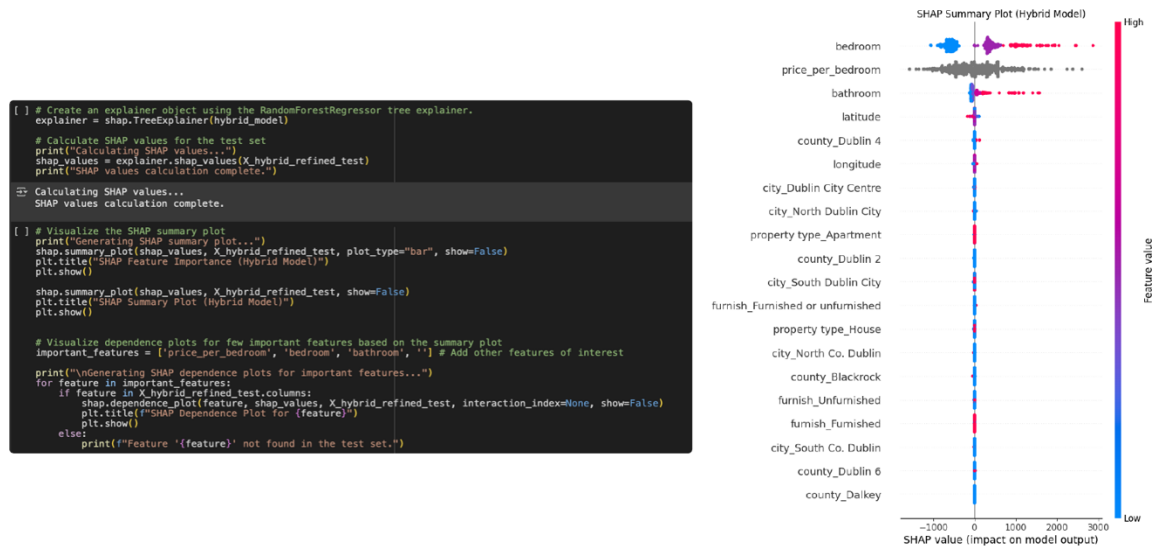


Figure 24: SHAP Implementation and Results

```
# Get the representative sample from the hybrid test data
feature_names = X_hybrid_refined_test.columns.tolist()

# Create a LIME explainer
explainer_lime = lime.lime_tabular.LimeTabularExplainer(
    training_data=X_hybrid_refined_test.values, # Use the training data values (numpy array)
    feature_names=feature_names,
    mode="regression"
)

# Choose an instance from the test set to explain (e.g., the first instance)
instance_idx = 0
instance_to_explain = X_hybrid_refined_test.iloc[instance_idx].values # Get the values as a numpy array
true_price = y_hybrid_refined_test.iloc[instance_idx] # Get the true price
predicted_price = hybrid_model.predict(instance_to_explain.reshape(1, -1))[0] # Predict price for the instance

print(f"Explaining prediction for instance {instance_idx}:")
print(f"True Price: {true_price:.2f}")
print(f"Predicted Price: {predicted_price:.2f}")

# Generate an explanation for the chosen instance
# num_features: number of features to include in the explanation
explanation = explainer_lime.explain_instance(
    data_row=instance_to_explain,
    predict_fn=hybrid_model.predict,
    num_features=10 # You can adjust the number of features
)

# Display the explanation
print("\nLIME Explanation:")
explanation.as_list()

[ ] # Visualize the LIME explanation
# This will typically open in a new browser tab or render in the notebook output
print("\nVisualizing LIME explanation:")
explanation.show_in_notebook(show_table=True, show_all=False)
```

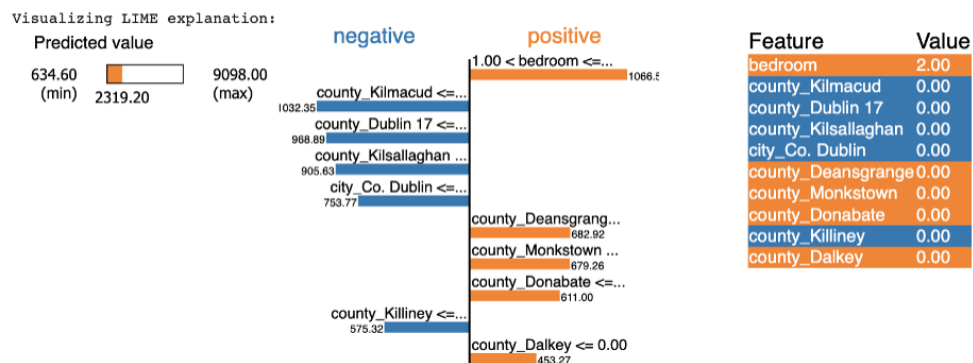


Figure 25: LIME Implementation and Results

6 Conclusion

This document presets a step-by-step approach in representing the technical implementation of this research. To replicate the code and achieve the same results and gain a deeper understanding for the inner working of the project, follow this outlined pathway.

[Report Presentation video link](#)

[Code Presentation video link](#)